

# verl 周报 2026 第 31 周 (07/27-08/02) : DeepSeek-V4 适配与 vLLM/Megatron 链路重构

verl-project/verl

周期: 2026-07-27 至 2026-08-02

来源 PR: 26 • 重点 PR: 24 • 自动生成

原文链接: <http://prhub.com.cn/verl-project/verl/reports/2026-07-27-to-2026-08-02>

## 执行摘要

本周共合并 26 个 PR (重点分析 24 个), 整体围绕两条主线展开: 一是 DeepSeek-V4 的支持链落地, 包括 vLLM 0.24.0/Megatron core\_v0.18.0 升级 (#7101)、FP8/MXFP4 量化权重同步修复与模块拆分 (#7224)、Megatron 后端新增 contiguous 上下文并行布局 (#7221); 二是权重同步与异步基础设施的集中治理, delta 分片协议在 Megatron-Bridge 和 VeOmni 后端打通 (#7181/#7085), TransferQueue 从 FD 泄漏修补走向全局事件循环重构 (#7157/#7162)。vLLM 侧同时清退了 0.18.0 以下兼容代码 (#7190), 并重构了权重同步控制流 (#7179)。整体看, 本周变化密度高但方向一致: 为下一代模型适配铺路, 同时偿还 rollout 恢复、trace 可观测性和资源泄漏方面的技术债。

## 本周重点变化

DeepSeek-V4 支持链成型。容器和 CI 层面, Dockerfile 改为 pip 预编译安装 vLLM 0.24.0、Megatron 桥接库升到 core\_v0.18.0, 并去掉 CI 运行时安装步骤 (#7101)。量化路径上, #7224 修复了 FP8 refit 后 `weight_scale_inv` 全零的静默错误, 根因是 bound method 被复制遮蔽类方法; 同时按 FP8/MXFP4 拆分 `vllm_fp8_utils.py` 与新增 `vllm_fp4_utils.py`, 并用 staging buffer 规避 CUDA graph 指针失效。Megatron 后端则新增 contiguous 上下文并行布局 (#7221), 对齐因子按布局自适应, 默认仍为 zigzag, 行为兼容。

delta 分片权重同步成为跨后端趋势。#7181 让 mcore 引擎加入 delta\_sharded 协议, 核心创新是 comm-stubbed probe: 保留真实进程组规模、用 NaN 哨兵穿过真实 `megatron_to_hf` 转换代码来推导 HF 坐标, 摆脱手工推导分片几何; #7085 则在 VeOmni 后端实现 EP 感知的 delta 导出, 通过自动 slot 表枚举让每个 rank 只导出本地变更块, 在 Qwen3-235B-A22B (ep8xfsd8) 上获得约 21x 加速。两个 PR 共用同一套协议边界 (`spec.py` 的 ShardSpec + HF 坐标 entry), 后端无关的协议设计开始显现。

vLLM 权重同步链路被大幅清理。#7190 将 vLLM 最低版本从 0.7.0 提到 0.18.0, 删掉 421 行旧版兼容分支; #7179 把 `update_weights_from_ipc` 重写为准备、接收、后处理三阶段, 并把量化入口统一为模块级 API; #7189 顺手移除无消费者的 `get_device_uuid`。这些改动降低了维护成本, 但也意味着对旧环境的破坏性变更, 需要紧随其后的端到端验证。

rollout 稳定性和可观测性补强。#7207 修复了 partial rollout resume 时 token 超预算的问题 (恢复路径的扣减逻辑因公式错误而未生效); #7158 将 `routed_experts` 合并从 `torch.cat` 改为 `np.concatenate`, 并让 SGLang `skip_tokenizer_init` 分支显式返回 numpy 数

组；#7204 重构了 trace 解码器，使 per-turn LLM generate 的 trace 也能输出 `prompt_text/response_text`；#7184 在 IPC 清理路径加 `is_support_ipc` 守卫，避免非 CUDA 设备崩溃。

## 模块与主题趋势

vLLM/rollout是本周热度最高的模块（10 个 PR 带 vllm 标签），热点文件集中在 `vllm_fp8_utils.py`、`vllm_rollout/utils.py` 与 `vllm_rollout.py`。趋势是：配合 vLLM 大版本升级，把量化权重同步、MoE 参数展开等逻辑从 rollout 层下沉 / 收敛到专用工具模块，同时为 DeepSeek-V4 的 FP8/MXFP4 双方案铺路。

Megatron/ 检查点呈现“协议统一”的趋势。`delta_sharded` 分片 delta 协议在 mcore 与 VeOmni 两个引擎落地，配套的 `transformer_impl.py`、`delta_checkpoint_engine.py` 成为高频文件；同时 Megatron 新增 Muon 优化器支持（#7120，内存 -18.3%、更新速度约 2x），以及模型合并器输出验证（#7193），说明训练后端在朝着更可组合、可验证的方向演进。

NPU/Ascend从零散补丁走向体系化：HCCL checkpoint 引擎支持大权重分块（#7205）、AscendDocker标签重构（#7201）、nightlyCI修复（#7176）以及多机启动文档（#7175），表明 NPU 作为一等公民的投入在持续增加，但 CI 覆盖和端到端验证仍是明显短板。

异步基础设施经历了一次典型的“临时修复→结构性重构”演进：#7157 补上临时事件循环的 `close()` 修复 FD 泄漏，#7162 随即引入全局共享事件循环和后台线程，并把 TransferQueue 配置抽取为独立 YAML。这类底层资源管理问题影响面大，值得继续关注全局循环的竞态与退出清理。

## 风险观察

本周风险集中在“核心路径变更 + 验证不足”的组合上。DeepSeek-V4 相关 PR 普遍标注缺少 GPU 端到端验证（#7221、#7207），而 #7101 的依赖升级又是大版本跨越；vLLM 权重同步链路同时经历删除旧代码（#7190）、控制流重写（#7179）和量化 bug 修复（#7224），三重变更叠加后，仅靠 CPU 单测不足以证明行为等价。

另一个需要紧盯的是 #7173：Qwen3.5 LoRA 与 MTP 支持被整体回滚，PR 与 issue 中社区成员两次询问原因均未获答复。即使回滚是出于稳定性考虑，也应尽快公开说明，否则会影响社区对 Megatron-Bridge 路线的信心。

NPU 侧，HCCL 分块传输依赖线程池长期占用（#7205 风险标注），且测试只在 NPU 上运行、CI 覆盖有限；Ascend 多机文档刚落地，尚未看到配套的自动化验证。TransferQueue 全局事件循环（#7162）也引入了全局共享状态，atexit 清理顺序与竞态条件需要在实际训练中观察。

## 重点 PR 速览

| PR        | 标题                                                                                     | 影响与看点                                                                     |
|-----------|----------------------------------------------------------------------------------------|---------------------------------------------------------------------------|
| #72<br>24 | [vllm] feat:<br>enhance<br>DeepSeek<br>V4 fp8/fp4<br>linear and<br>moe weight<br>refit | 修复 FP8 refit 静默错误, 拆分 FP8/MXFP4 同步逻辑; staging buffer 与 CUDA graph 交互是核心设计 |
| #71<br>81 | [megatron]<br>feat: delta_<br>sharded on<br>Megatron-B<br>ridge param<br>mappings      | mcore 接入 delta_sharded; NaN 探针推导 HF 坐标, 规避手工索引; PP 显式报错兜底                 |
| #70<br>85 | [veomni]<br>feat: EP-aw<br>are sharded<br>delta export                                 | VeOmni EP+FSDP 场景 delta 导出, Qwen3-235B 约 21x 加速; slot 表自动枚举避免手动维护         |
| #71<br>01 | [docker]<br>feat: upgrade vllm and<br>megatron<br>version                              | vLLM 0.24.0/Megatron core_v0.18.0 升级, 适配 MoE 工厂函数重构, CI 移除运行时安装           |
| #72<br>07 | [rollout]<br>fix: cap<br>partial rollout<br>resumes                                    | 修复恢复时 token 超发, 恢复预算扣减逻辑生效; CPU 测试用 fake server 复刻公式                      |
| #72<br>04 | [rollout]<br>fix: decode<br>per-turn<br>LLM tokens<br>in traces                        | trace 解码器支持 per-turn generate; 采用窄范围修复而非通用 field-mapping, 避免新增公共 API      |
| #71<br>62 | [perf] fix:<br>Prevent<br>creating<br>new threads/<br>event loop                       | 全局共享事件循环替代每次调用创建, 抽取 TransferQueue 独立 YAML, 消除 FD 泄漏隐患                    |

| PR    | 标题                                | 影响与看点                                   |
|-------|-----------------------------------|-----------------------------------------|
| #7173 | Revert Qwen3.5 LoRA & MTP support | 回滚 #5599，删除 8 个权重名映射函数；回滚原因未公开，社区有疑问未答复 |

除了上述重点 PR，#7120（Muon 优化器）展示了 optimizer 接入的完整范式，[muon\\_match\\_adamw\\_update\\_rms](#) 的自动缩放因子推导值得借鉴；#7193 的模型合并器输出验证、#7188 的解耦 PPO 支持也分别在模型导出和训练算法路径上补齐了关键能力。

### 后续建议

1. 尽快补齐 DeepSeek-V4 端到端验证。建议以 #7101 的新镜像为基线，跑一轮包含 Megatron 训练 + vLLM rollout + FP8/MXFP4 权重同步的完整训练，重点验证 #7221 的 contiguous CP 布局与 #7224 的 staging buffer 在真实 CUDA graph 场景下的正确性。
2. 为 vLLM 权重同步链路增加 GPU 集成测试。#7179 与 #7224 都属于“重构 + 修 bug”叠加，当前主要依赖 CPU 单测和人工 review；建议在 CI 中增加至少一个 FP8 Qwen/DeepSeek 模型的 rollout 权重同步用例。
3. 公开 #7173 的回滚原因与后续路线。无论是临时回退还是放弃该能力，都应在 issue 中明确回应社区，并给出 Qwen3.5 LoRA/MTP 的重新支持时间表。
4. NPU 路径验证前置。HCCL 分块传输与 Ascend CI 的改动建议在 NPU 集群上做多机、多卡的压测，并评估线程池长期占用对训练吞吐的影响；TransferQueue 全局事件循环上线后，需关注长时间运行下的稳定性与退出清理。
5. 持续推广 delta 分片协议。本周 mcore 与 VeOmni 已统一到同一套协议边界，建议后续将 FSDP 后端的权重同步也纳入 delta\_sharded 体系，并沉淀协议文档，降低多后端维护成本。